
Oral presentation | Data science and AI

Data science and AI-III

Thu. Jul 18, 2024 2:00 PM - 4:00 PM Room C

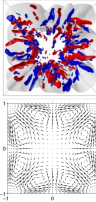
[11-C-02] Data-driven thermal control of secondary flow in a marginally turbulent square-duct flow and relevance to the invariant solutions

*Atsushi Sekimoto¹, Takashi Mitani¹ (1. Okayama University)

Keywords: Turbulent square-duct flow , data-driven control, optimization

2024.7.14 -19 ICCFD 12

Data-driven thermal control of secondary flow in a marginally turbulent square-duct and relevance to the invariant solutions



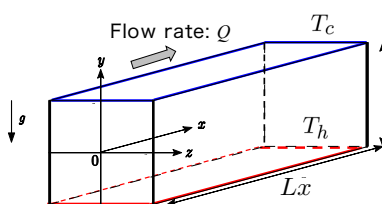
○ Atsushi Sekimoto
Okayama University
Faculty of Environmental, Life, and Natural Science,
Mathematical data sciences, Assoc. Prof.

Takashi Mitani
Okayama University, graduate school of
Environmental, Life, and Natural Science

※ Most of computation was performed at "Wisteria" in U. of Tokyo

1

Geometry & simple heating from below



Flow rate: Q

T_c

T_h

H

Lx

$-h$

g

$u_b = \frac{Q}{H^2}$

$\Delta T = T_h - T_c$

$T_0 = \frac{T_h + T_c}{2}$

Governing equation

$$\frac{\partial \mathbf{u}}{\partial t} + (\mathbf{u} \cdot \nabla) \mathbf{u} = -\nabla \left(\frac{p}{\rho} \right) + \frac{1}{Re_H} \nabla^2 \mathbf{u} + \frac{Gr}{Re_H^2} T^* \mathbf{e}_y$$

$$\frac{\partial T^*}{\partial t} + (\mathbf{u} \cdot \nabla) T^* = \frac{1}{Pr Re_H} \nabla^2 T^*$$

Ri

2

Numerical method

- Time:** pressure correction fractional-step method semi-implicit 3-step Runge-Kutta [Verzicco & Orlandi, *J. Comput. Phys.* (1996)]
CFL < 0.3
- Space:** Pseudo-spectral method streamwise(x), Fourier expansions cross-streamwise(y, z), Chebyshev polynomials
 $\Delta x^+ < 16.4$ $\Delta y^+, \Delta z^+ < 5.4$
- 2D Helmholtz equations:** Chebyshev collocation (fast diagonalization technique) [Haldenwang, *J. Comput. Phys.* (1984)]
⇒ "dgemm"×4 → GPU acceleration for high-Re but low-Re in my talk, today

3

Coherent Structures in a minimal box

$Q = -\Delta p = 0.03 u_\tau^4 / \nu^2$

Vortices: $Re_H = 2200$

$Re_H = 4400$

Streak $u = 0.6 u_b$

$H^+ = 155$

$H^+ = 300$

$H^+ = H / (u_\tau / \nu)$

$u_\tau \equiv \sqrt{\frac{\tau_w}{\rho}}$

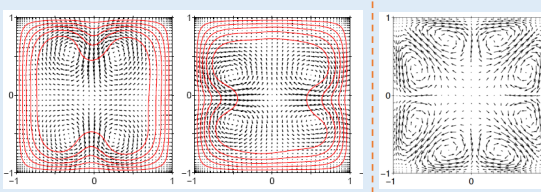
τ_w : wall shear stress

4

Secondary flows of Prandtl's second kind at marginal-Re turbulence $Ri = 0$

$Re_H = 2200$

$Re_H = 4400$



4-vortex

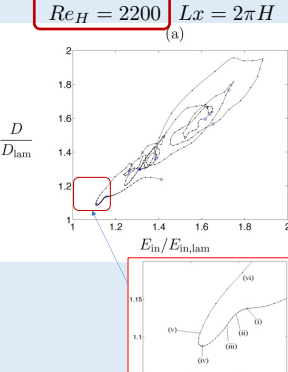
Uhlmann, Pinelli, Kawahara & Sekimoto, *J. Fluid Mech.* (2007)

5

Dissipation vs Energy input $Ri = 0$

$Re_H = 2200$ $Lx = 2\pi H$

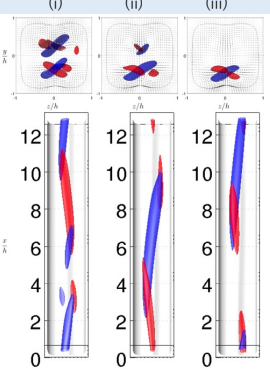
(a)



$\frac{D}{D_{lam}}$

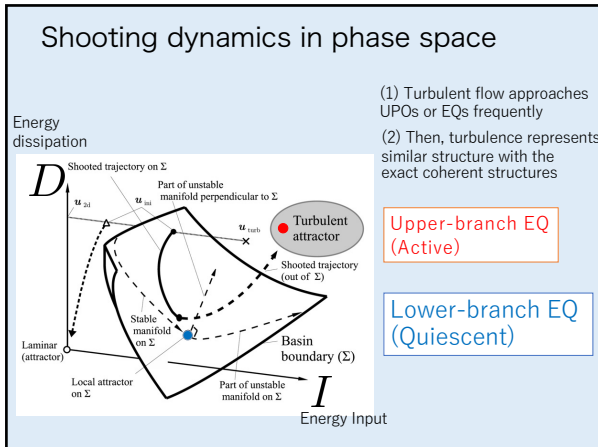
$\frac{E_{in}}{E_{in,lam}}$

(i) (ii) (iii)

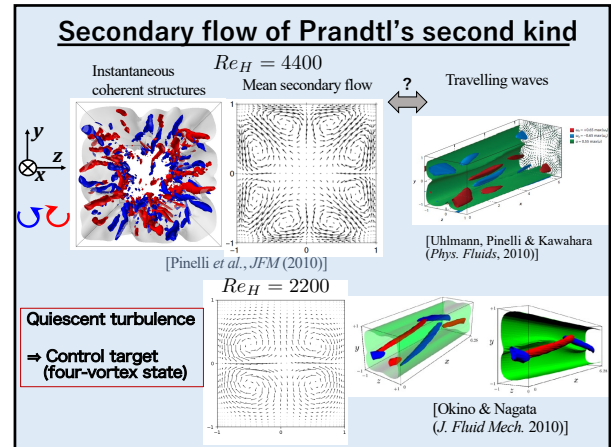


z/h

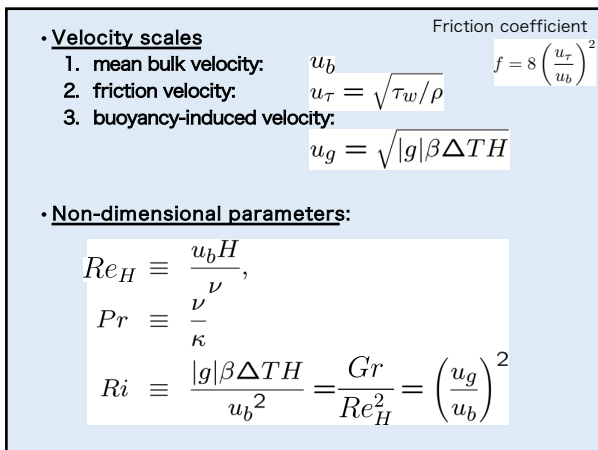
6



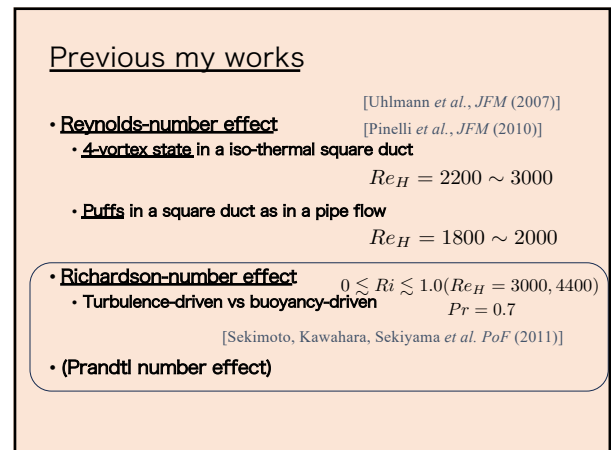
7



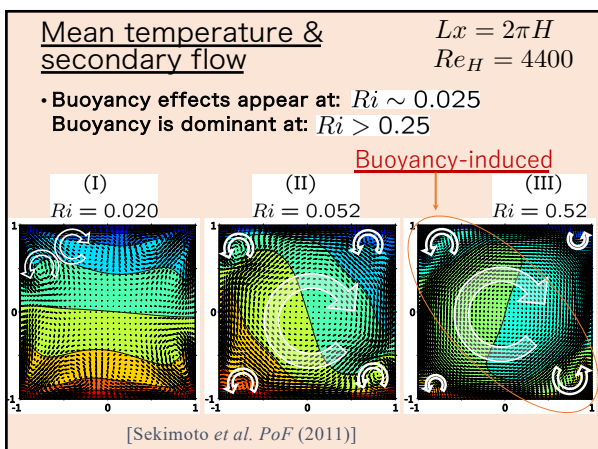
8



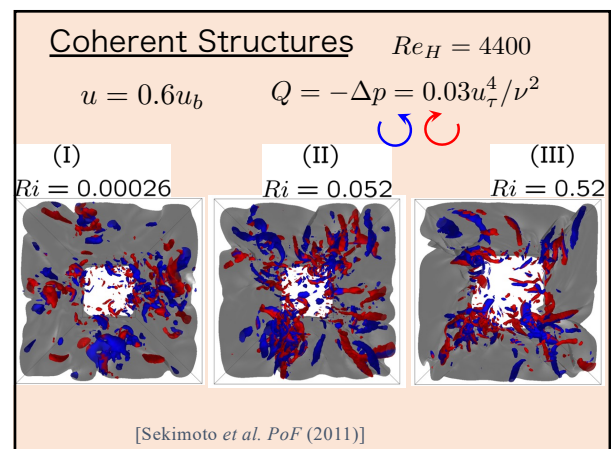
9



10



11



12

Stabilization of a 4-vortex state $Pr = 0.7$ by heating from below $Lx = 20\pi H$

	$Ri = 0$	$Ri = 0.01$	$Ri = 0.1$
$Re_H = 2200$	Turb.	Turb.	4-vortex
$Re_H = 2000$	Puff	Puff	4-vortex
$Re_H = 1800$	Puff	Puff	4-vortex
$Re_H = 1600$	Laminar	Laminar	Laminar

- Buoyancy stabilises a 4-vortex-type secondary flow at moderate Richardson numbers, suggesting a control strategy by using buoyancy force

- Automatic complex control
→ deep reinforcement learning



13

Deep Reinforcement Learning

- DRL, maximize the expectation, $E[R_L]$ of total reward R_L to find the optimal policy

$$\Pi^* = \underset{\Pi}{\operatorname{argmax}} \{E[R_L]\}$$

$$R_L = \sum_{l=1}^T \gamma^l r_l$$

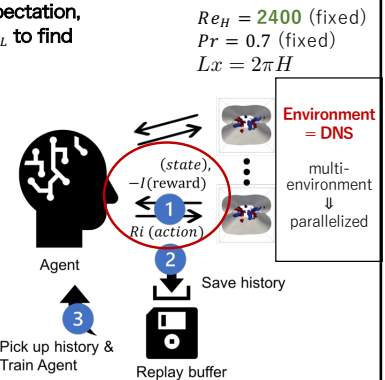
Π Policy = Control of Ri

Π^* Optimal policy (control)

R_L Total reward accumulated = Cost function

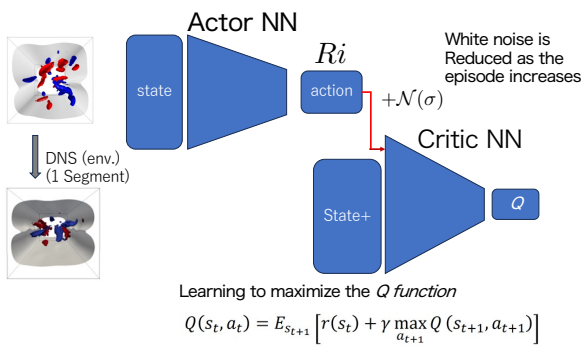
γ Reduction rate = weight of past reward

r_l Reward at each segment (step)



14

Framework of DDPC algorithm



15

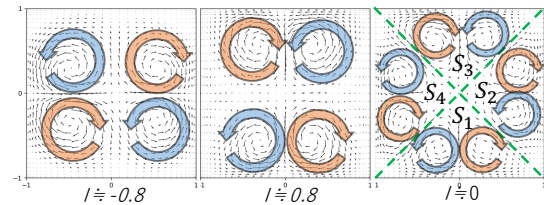
Reward: four-vortex state indicator, I

- The indicator function using the mean streamwise vorticity [Uhlmann *et al.* (2007)]

$$I(t) \equiv \frac{S_1 + S_3 - S_2 - S_4}{S_1 + S_3 + S_2 + S_4}, \quad S_i = \iint_{\Omega_i} \langle \omega_x \rangle_x^2 dy dz$$

$$\Omega_1: \{(y, z) | y < z \wedge y < -z\}, \Omega_2: \{(y, z) | y < z \wedge y > -z\},$$

$$\Omega_3: \{(y, z) | y > z \wedge y > -z\}, \Omega_4: \{(y, z) | y > z \wedge y < -z\}$$



16

Cost function in RL

- $Re_H = 2400$, $Pr = 0.7$ (fixed), control the heat input $Ri = 0.0 - 0.1$ (limited range) observing environment (DNS)



17

Details of RL

- Algorithm

- DDPG [Lillicrap *et al.* (2016)]

- Two DNNs: ADAM optimizer

- State [= what we observe from env. (DNS)]

- The Fourier modes of cross-sectional velocities, v, w [0, 1, 2] → Focusing on Large-scale motion

- Action (control) $Ri = [0 \sim 0.1]$

$Re_H = 2400$ (fixed)

$Pr = 0.7$ (fixed)

$Lx = 2\pi H = 4\pi h$

- Reward [from env. (DNS)]

- based on the indicator function

If $I > 0.5$

Reward, $r = 0$

Else

Reward, $r = -I$

- Episode (=1 run of DNS)

- 1 episode $\approx 5400 h/u_b$
- Divide each episode into 200 of segments
- The different initial condition for each episode (30 episode)

- Segment (the interval of control)

- 1 segment $= 27 \frac{h}{u_b}$
- no history correlations between previous control (~2 washout time) → "Markovian"

18

The structure of Neural Networks 🧠

- DDPG has two NN: Actor NN calculates the action from the state, and the Critic NN does Q of actions for each state
- Actor NN
 - $\text{state}(13068) \Rightarrow (32768) \Rightarrow (4096) \Rightarrow (512) \Rightarrow (64) \Rightarrow \text{action}$
- Critic NN
 - $\text{state+}, \text{action}(13069) \Rightarrow (32768) \Rightarrow (4096) \Rightarrow (512) \Rightarrow (64) \Rightarrow \text{Q-value}$
- Activation function: ReLU
- Fully-connected

19

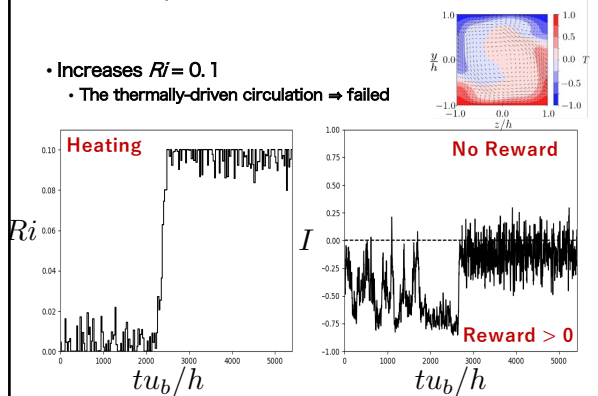
Hyper parameter for NNs

- Learning rate
 - Actor_lr = 0.0001
 - Critic_lr = 0.001
- Reduction rate (almost no consideration from past rewards)
 - 0.99
- soft update parameter (target NN)
 - 0.005
- Replay buffer (common for all episodes and agents, necessary for parallel environments)
 - 2000 buffer (remember the latest 10 episodes)
 - Butch size: 64

20

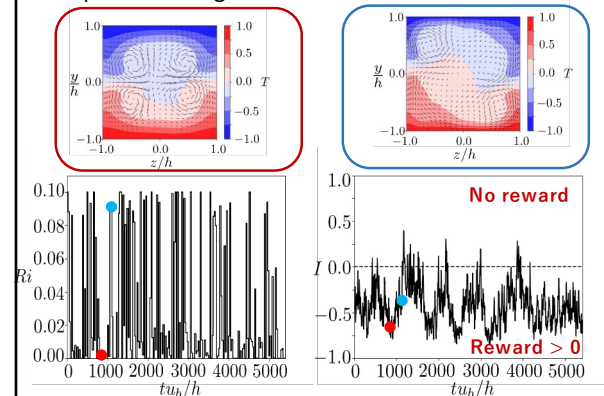
Result: episode1

- Increases $Ri = 0.1$
- The thermally-driven circulation $\Rightarrow$ failed



21

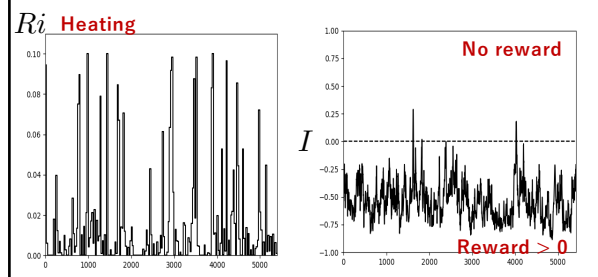
Episode 3, agent becomes smarter



22

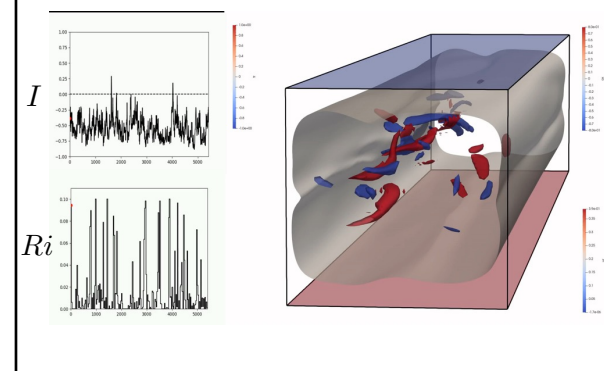
Episode 8

- Switching between $Ri = 0$ and $Ri = 0.1$
- An intermittent heat control is realized
- It keeps one on the secondary flow pattern of $I < 0$



23

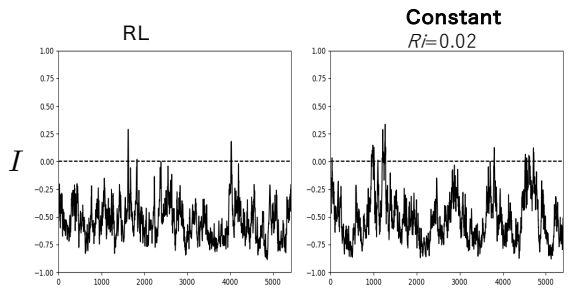
Control: ON (Movie)



24

Constant heating ($Ri=0.02$) vs RL-control
[from the same initial condition]

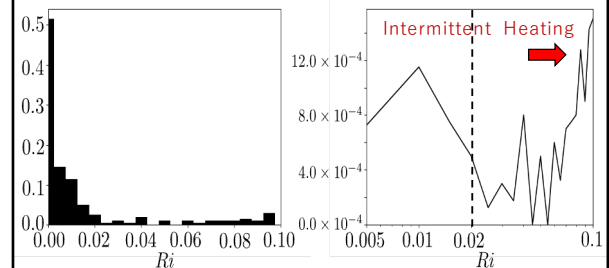
- Dynamic heating controlled by RL achieves more-stable four-vortex pattern...



25

Contribution of intermittent heating

- $Ri \rightarrow$ probability density function (PDF, left)
→ pre-multiplied Ri , PDF is multiplied by Ri (right)
- In logarithmic scale in $Ri \rightarrow$ Area corresponds to the total heat input
- Compare to the constant heating ($Ri=0.02$) $\rightarrow$ 20% heat reduction



26

Summary and perspectives

- Complex 3-D flow with Re -, Ri -, Pe numbers
 - (also aspect ratios, As , Computational domain, Lx , etc.)
 - Stabilized the mean secondary flow pattern, which is called as a corner flow, using buoyancy force
→ heat enhancement, particle migrations (or separations), control in a micro-reaction channel
- Difficulties of flow control from the strong nonlinearity of turbulence
 - Successful nonlinear control method using (deep) reinforcement learning
→ now, targeting a specific known flow pattern (4-vortex state)
 - Background physics is important for the further application with ML techniques
→ AI technology
 - The target flow is realizable? Dynamical system approach of turbulence
 - Which ML models are suitable? How do we reduce the complexity $\rightarrow$ GPU mem.
 - How about high- Re ?
- High-fidelity CFD are expensive
 - Parameter optimization using Bayes optimization or GA algorithm
 - Keep developing the advanced GPU-accelerated code (or QPU-accelerated?)
 - Scale up
→ Power-game in industry and international-collaboration in academic

27

Thank you very much

Questions and discussions $\rightarrow$ e-mail: asekimoto@okayama-u.ac.jp
Please visit: AtsushiSekimotoLab.
(Data-driven computing, dynamical system, flow control)
<https://sites.google.com/view/sekimoto-lab/research>

Acknowledgements (All in Japanese)

- 九州大学情報基盤研究開発センター, 先端的計算科学研究プロジェクト (2019ベンチマーク課題)
- 東京大学情報基盤センター若手利用課題 (2019, 2020, 2022)
(M2三谷: 2023年度若手利用課題)
- ホソカワ粉体工学財団 研究助成 (2022年度)
- 稲盛財団 研究助成 (2023年度)
- 九州大学マスコアインダストリ 共同利用研究 2023
代表: 中澤崇 (阪大MMDS)

28